

The 10 Building Blocks That Separate Working LLM Systems from Broken Ones



Forget prompt tricks — the real engineering happens in the infrastructure around the model

If you want to understand why some AI applications feel polished and reliable while others fall apart the moment you push them beyond a demo, the answer almost never lies in the prompt. The real difference is architectural. Modern LLM applications are not simply a question and a response. They are layered systems that manage context windows, orchestrate external tools, retrieve live data, and coordinate multi-step workflows behind every interaction. The engineers building these systems spend far more time on plumbing than on prose. Here are ten foundational concepts that define how serious LLM systems are actually constructed — and why each one matters.

01

Context Engineering Has Replaced Prompt Engineering

The industry has quietly moved beyond the era of clever prompt tricks. What matters now is context engineering — the discipline of deciding what information reaches the model, in what order, and in what format. This encompasses system instructions, conversation history, retrieved documents, tool definitions, memory, intermediate reasoning steps, and execution traces. A well-engineered context window can make a mediocre prompt perform brilliantly, while a poorly managed one will sabotage even the most carefully worded instruction. Most LLM failures in production are not caused by bad prompts. They happen because the context feeding the model is missing critical information, contains outdated data, is cluttered with irrelevant noise, or presents facts in an order that confuses the model's reasoning. Mastering context management is the single highest-leverage skill in LLM engineering today.

02

Tool Calling Connects Models to the Real World

A language model trained on static data will eventually give you a confident but wrong answer about something that changed last week. Tool calling solves this by allowing the model to invoke external functions instead of relying solely on its training data. In practice, this is how an LLM searches the web, queries a database, executes code, calls an API, or retrieves a document. The model decides which tool to use, generates the appropriate input parameters, and incorporates the result back into its response. This mechanism transforms LLMs from sophisticated text predictors into systems that can interact with live information and take real actions — the foundation of every useful AI application in production today.

For more content like and follow me:



@bertblevins

INCGPT.COM

03

The Model Context Protocol Creates a Universal Plug-and-Play Standard

Tool calling works, but every integration used to require custom code. The Model Context Protocol, or MCP, changes that by establishing a standardized way to expose tools, data sources, and services so that any compatible AI assistant can use them. Think of it as a universal adapter. Instead of building bespoke connectors for every combination of model and service, developers expose their capabilities through a consistent interface. Any LLM application that speaks MCP can immediately discover and use those capabilities without custom wiring. This dramatically reduces integration effort and makes it practical to build AI systems that connect to dozens of enterprise tools without drowning in maintenance overhead.

04

Retrieval-Augmented Generation Grounds Answers in Facts

RAG is the technique that keeps LLMs honest. Instead of relying on the model's baked-in knowledge, a RAG system retrieves relevant documents from an external knowledge base and inserts them into the context before the model generates a response. The model then answers based on what it has been shown, not what it vaguely remembers from training. This approach is essential for any application where accuracy matters: legal research, medical information, technical documentation, or customer support. The quality of a RAG system depends heavily on how documents are chunked, how embeddings are generated, how retrieval is ranked, and how the retrieved content is formatted before it reaches the model. Getting these details right is what separates a useful knowledge assistant from one that hallucinates plausibly.

05

Agentic Loops Let Models Plan and Execute Multi-Step Tasks

A standard LLM interaction is a single turn: question in, answer out. An agentic loop extends this into a cycle where the model can reason about a problem, take an action, observe the result, and decide what to do next. This repeating loop — think, act, observe, repeat — is what allows AI systems to tackle complex goals that require multiple steps, such as researching a topic across several sources, debugging code iteratively, or managing a multi-stage workflow. The engineering challenge lies in controlling the loop: preventing infinite cycles, managing costs from repeated model calls, handling errors at each step, and knowing when the task is genuinely complete. Poorly designed loops are the most common source of runaway costs and unreliable behavior in agentic AI applications.

06

Agent-to-Agent Communication Enables Team-Based AI

Many real-world tasks are too complex for a single agent. A research task might require one agent to gather information, another to analyze it, and a third to draft a report. Agent-to-agent communication protocols — like Google's A2A standard — provide a structured way for these specialized agents to coordinate. Each agent handles its domain of expertise and passes results to the next in the chain. This mirrors how human teams operate: a specialist hands off to another specialist, with clear interfaces between roles. The protocol ensures these handoffs are secure, structured, and auditable — critical requirements for enterprise environments where compliance and accountability matter.

07

Structured Outputs Guarantee Machine-Readable Results

LLMs are natural language machines, but downstream systems speak JSON, SQL, and structured data. Structured outputs bridge this gap by constraining the model to produce responses in a specific, predictable format. Instead of hoping the model returns valid JSON and parsing whatever comes back, engineers can now enforce a schema so that the output is guaranteed to match the expected structure. This is essential for any system where the LLM's output feeds directly into another process — populating a database, triggering an API call, or filling a form. Without structured outputs, LLM pipelines require fragile parsing logic that breaks whenever the model decides to phrase something slightly differently.

08

Guardrails and Evaluation Keep Systems Safe and Honest

Deploying an LLM without guardrails is like launching a website without input validation. Guardrails are validation layers that check both the input going into the model and the output coming out. They can block harmful or off-topic requests, detect hallucinated facts, enforce policy compliance, and ensure the model stays within its intended scope. Evaluation complements guardrails by providing systematic measurement of how well the system performs. This includes automated benchmarks, human review workflows, and the increasingly common 'LLM-as-a-judge' approach, where a separate model evaluates the quality of another model's outputs. Together, guardrails and evaluation form the safety net that makes the difference between a prototype and a production system.



09

Prompt Caching Slashes Costs and Latency

Every token sent to an LLM costs money and adds latency. In many applications, large portions of the prompt — system instructions, tool definitions, reference documents — remain identical across thousands of requests. Prompt caching allows these stable portions to be processed once and reused, rather than re-sent and re-computed every time. The practical impact is significant: applications that process the same system prompt and document set repeatedly can see cost reductions of 50–90% on the cached portion and noticeably faster response times. For high-volume production systems, caching is not an optimization — it is a requirement for economic viability.

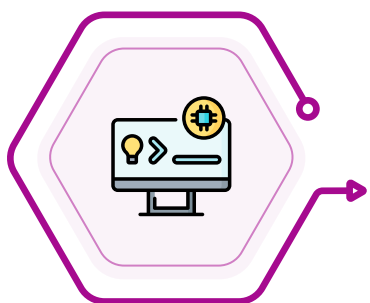
10

Observability and Tracing Make the Invisible Visible

LLM systems are notoriously difficult to debug because so much happens inside a black box. Observability tools address this by logging every step of the pipeline: what context was assembled, which tools were called, what the model returned, how long each step took, and how much it cost. Tracing connects these individual events into a coherent story for each request. When a user reports a bad answer, engineers can trace back through the entire chain — from the initial query through retrieval, context assembly, model inference, and post-processing — to pinpoint exactly where things went wrong. Without observability, debugging LLM systems is guesswork. With it, every failure becomes a learning opportunity, and every improvement can be verified with data

The Shift That Matters

The field of LLM engineering has matured rapidly. The conversation has moved from “how do I write a better prompt?” to “how do I build a system that manages context, coordinates agents, retrieves live data, enforces safety, and scales economically?” These ten building blocks are the vocabulary of that new conversation. The engineers who understand them are not just building chatbots. They are building the infrastructure layer of a new generation of intelligent software — and that shift from prompt craft to systems engineering is where the real opportunity lies.



The evolution of Large Language Model (LLM) systems highlights a significant shift from basic prompt engineering to complex systems engineering. The key concepts discussed, such as context engineering, tool calling, and retrieval-augmented generation, emphasize the importance of managing data, tools, and processes effectively to create efficient and reliable AI systems. Modern LLMs now extend beyond simple question-answer tasks to execute multi-step workflows, interact with live data, and ensure output integrity through structured formats and safety measures. The integration of agent-based systems and observability tools further enhances these models' ability to handle complex tasks with precision and scalability. As LLM systems mature, the real opportunity lies in mastering these foundational building blocks, which will enable the development of intelligent applications that are both powerful and resilient in real-world production environments.

